

CHAPTER 7

REVEALED PREFERENCE

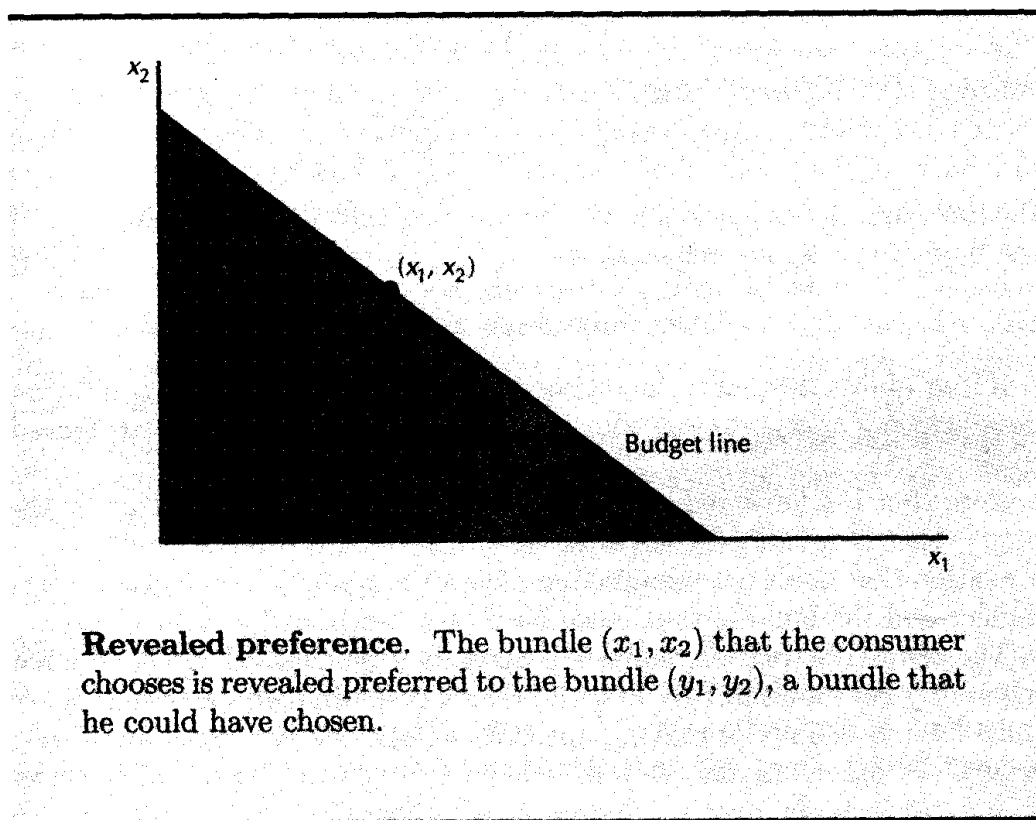
In Chapter 6 we saw how we can use information about the consumer's preferences and budget constraint to determine his or her demand. In this chapter we reverse this process and show how we can use information about the consumer's demand to discover information about his or her preferences. Up until now, we were thinking about what preferences could tell us about people's behavior. But in real life, preferences are not directly observable: we have to discover people's preferences from observing their behavior. In this chapter we'll develop some tools to do this.

When we talk of determining people's preferences from observing their behavior, we have to assume that the preferences will remain unchanged while we observe the behavior. Over very long time spans, this is not very reasonable. But for the monthly or quarterly time spans that economists usually deal with, it seems unlikely that a particular consumer's tastes would change radically. Thus we will adopt a maintained hypothesis that the consumer's preferences are stable over the time period for which we observe his or her choice behavior.

7.1 The Idea of Revealed Preference

Before we begin this investigation, let's adopt the convention that in this chapter, the underlying preferences—whatever they may be—are known to be strictly convex. Thus there will be a *unique* demanded bundle at each budget. This assumption is not necessary for the theory of revealed preference, but the exposition will be simpler with it.

Consider Figure 7.1, where we have depicted a consumer's demanded bundle, (x_1, x_2) , and another arbitrary bundle, (y_1, y_2) , that is beneath the consumer's budget line. Suppose that we are willing to postulate that this consumer is an optimizing consumer of the sort we have been studying. What can we say about the consumer's preferences between these two bundles of goods?



Well, the bundle (y_1, y_2) is certainly an affordable purchase at the given budget—the consumer could have bought it if he or she wanted to, and would even have had money left over. Since (x_1, x_2) is the *optimal* bundle, it must be better than anything else that the consumer could afford. Hence, in particular it must be better than (y_1, y_2) .

The same argument holds for any bundle on or underneath the budget line other than the demanded bundle. Since it *could* have been bought at

the given budget but wasn't, then what *was* bought must be better. Here is where we use the assumption that there is a *unique* demanded bundle for each budget. If preferences are not strictly convex, so that indifference curves have flat spots, it may be that some bundles that are *on* the budget line might be just as good as the demanded bundle. This complication can be handled without too much difficulty, but it is easier to just assume it away.

In Figure 7.1 all of the bundles in the shaded area underneath the budget line are revealed worse than the demanded bundle (x_1, x_2) . This is because they could have been chosen, but were rejected in favor of (x_1, x_2) . We will now translate this geometric discussion of revealed preference into algebra.

Let (x_1, x_2) be the bundle purchased at prices (p_1, p_2) when the consumer has income m . What does it mean to say that (y_1, y_2) is affordable at those prices and income? It simply means that (y_1, y_2) satisfies the budget constraint

$$p_1 y_1 + p_2 y_2 \leq m.$$

Since (x_1, x_2) is actually bought at the given budget, it must satisfy the budget constraint with equality

$$p_1 x_1 + p_2 x_2 = m.$$

Putting these two equations together, the fact that (y_1, y_2) is affordable at the budget (p_1, p_2, m) means that

$$p_1 x_1 + p_2 x_2 \geq p_1 y_1 + p_2 y_2.$$

If the above inequality is satisfied and (y_1, y_2) is actually a different bundle from (x_1, x_2) , we say that (x_1, x_2) is **directly revealed preferred** to (y_1, y_2) .

Note that the left-hand side of this inequality is the expenditure on the bundle that is *actually chosen* at prices (p_1, p_2) . Thus revealed preference is a relation that holds between the bundle that is actually demanded at some budget and the bundles that *could have been* demanded at that budget.

The term “revealed preference” is actually a bit misleading. It does not inherently have anything to do with preferences, although we've seen above that if the consumer is making optimal choices, the two ideas are closely related. Instead of saying “ X is revealed preferred to Y ,” it would be better to say “ X is chosen over Y .” When we say that X is revealed preferred to Y , all we are claiming is that X is chosen when Y could have been chosen; that is, that $p_1 x_1 + p_2 x_2 \geq p_1 y_1 + p_2 y_2$.

7.2 From Revealed Preference to Preference

We can summarize the above section very simply. It follows from our model of consumer behavior—that people are choosing the best things they can

afford—that the choices they make are preferred to the choices that they could have made. Or, in the terminology of the last section, if (x_1, x_2) is *directly revealed preferred* to (y_1, y_2) , then (x_1, x_2) is in fact *preferred* to (y_1, y_2) . Let us state this principle more formally:

The Principle of Revealed Preference. *Let (x_1, x_2) be the chosen bundle when prices are (p_1, p_2) , and let (y_1, y_2) be some other bundle such that $p_1x_1 + p_2x_2 \geq p_1y_1 + p_2y_2$. Then if the consumer is choosing the most preferred bundle she can afford, we must have $(x_1, x_2) \succ (y_1, y_2)$.*

When you first encounter this principle, it may seem circular. If X is revealed preferred to Y , doesn't that automatically mean that X is preferred to Y ? The answer is no. "Revealed preferred" just means that X was chosen when Y was affordable; "preference" means that the consumer ranks X ahead of Y . If the consumer chooses the best bundles she can afford, then "revealed preference" implies "preference," but that is a consequence of the model of behavior, not the definitions of the terms.

This is why it would be better to say that one bundle is "chosen over" another, as suggested above. Then we would state the principle of revealed preference by saying: "If a bundle X is chosen over a bundle Y , then X must be preferred to Y ." In this statement it is clear how the model of behavior allows us to use observed choices to infer something about the underlying preferences.

Whatever terminology you use, the essential point is clear: if we observe that one bundle is chosen when another one is affordable, then we have learned something about the preferences between the two bundles: namely, that the first is preferred to the second.

Now suppose that we happen to know that (y_1, y_2) is a demanded bundle at prices (q_1, q_2) and that (y_1, y_2) is itself revealed preferred to some other bundle (z_1, z_2) . That is,

$$q_1y_1 + q_2y_2 \geq q_1z_1 + q_2z_2.$$

Then we know that $(x_1, x_2) \succ (y_1, y_2)$ and that $(y_1, y_2) \succ (z_1, z_2)$. From the transitivity assumption we can conclude that $(x_1, x_2) \succ (z_1, z_2)$.

This argument is illustrated in Figure 7.2. Revealed preference and transitivity tell us that (x_1, x_2) must be better than (z_1, z_2) for the consumer who made the illustrated choices.

It is natural to say that in this case (x_1, x_2) is **indirectly revealed preferred** to (z_1, z_2) . Of course the "chain" of observed choices may be longer than just three: if bundle A is directly revealed preferred to B , and B to C , and C to D , ... all the way to M , say, then bundle A is still indirectly revealed preferred to M . The chain of direct comparisons can be of any length.

If a bundle is either directly or indirectly revealed preferred to another bundle, we will say that the first bundle is **revealed preferred** to the

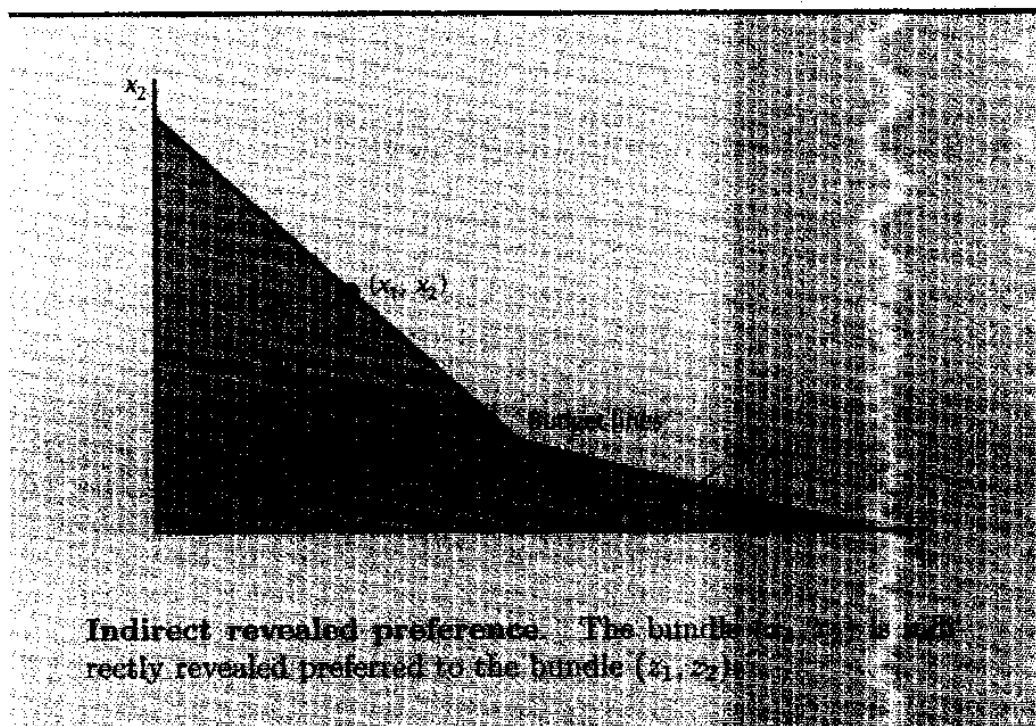


Figure
7.2

Indirect revealed preference. The bundle (x_1, x_2) is indirectly revealed preferred to the bundle (x_1, x_2) .

second. The idea of revealed preference is simple, but it is surprisingly powerful. Just looking at a consumer's choices can give us a lot of information about the underlying preferences. Consider, for example, Figure 7.2. Here we have several observations on demanded bundles at different budgets. We can conclude from these observations that since (x_1, x_2) is revealed preferred, either directly or indirectly, to all of the bundles in the shaded area, (x_1, x_2) is in fact *preferred* to those bundles by the consumer who made these choices. Another way to say this is to note that the true indifference curve through (x_1, x_2) , whatever it is, must lie above the shaded region.

7.3 Recovering Preferences

By observing choices made by the consumer, we can learn about his or her preferences. As we observe more and more choices, we can get a better and better estimate of what the consumer's preferences are like.

Such information about preferences can be very important in making policy decisions. Most economic policy involves trading off some goods for others: if we put a tax on shoes and subsidize clothing, we'll probably end up having more clothes and fewer shoes. In order to evaluate the desirability of such a policy, it is important to have some idea of what consumer preferences between clothes and shoes look like. By examining consumer choices, we can extract such information through the use of revealed preference and related techniques.

If we are willing to add more assumptions about consumer preferences, we can get more precise estimates about the shape of indifference curves. For example, suppose we observe two bundles Y and Z that are revealed preferred to X , as in Figure 7.3, and that we are willing to postulate preferences are convex. Then we know that all of the weighted averages of Y and Z are preferred to X as well. If we are willing to assume that preferences are monotonic, then all the bundles that have more of both goods than X , Y , and Z —or any of their weighted averages—are also preferred to X .

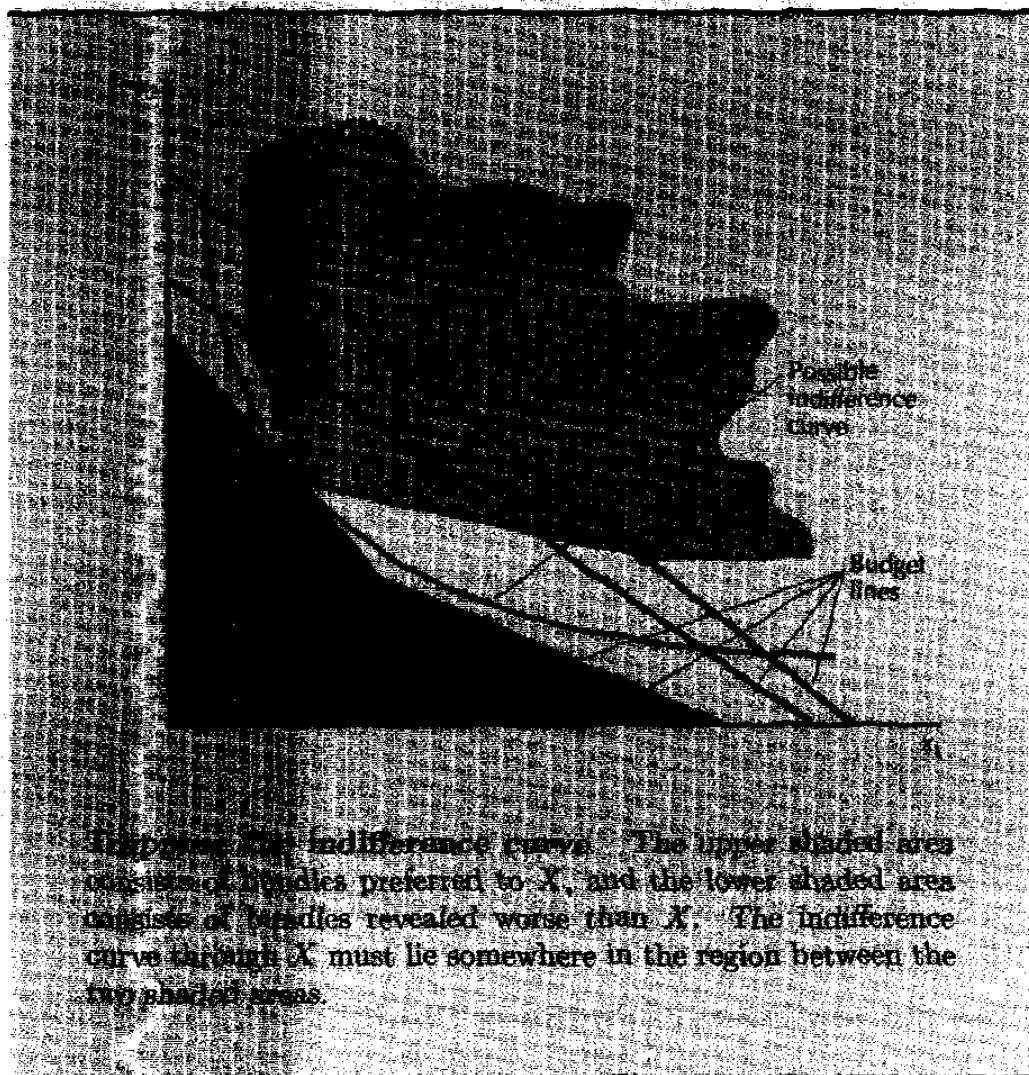


Figure
7.3

The region labeled "Worse bundles" in Figure 7.3 consists of all the bundles to which X is revealed preferred. That is, this region consists of all the bundles that cost less than X , along with all the bundles that cost less than bundles that cost less than X , and so on.

Thus, in Figure 7.3, we can conclude that all of the bundles in the upper shaded area are better than X , and that all of the bundles in the lower shaded area are worse than X , according to the preferences of the consumer who made the choices. The true indifference curve through X must lie somewhere between the two shaded sets. We've managed to trap the indifference curve quite tightly simply by an intelligent application of the idea of revealed preference and a few simple assumptions about preferences.

7.4 The Weak Axiom of Revealed Preference

All of the above relies on the assumption that the consumer *has* preferences and that she is always choosing the best bundle of goods she can afford. If the consumer is not behaving this way, the “estimates” of the indifference curves that we constructed above have no meaning. The question naturally arises: how can we tell if the consumer is following the maximizing model? Or, to turn it around: what kind of observation would lead us to conclude that the consumer was *not* maximizing?

Consider the situation illustrated in Figure 7.4. Could both of these choices be generated by a maximizing consumer? According to the logic of revealed preference, Figure 7.4 allows us to conclude two things: (1) (x_1, x_2) is preferred to (y_1, y_2) ; and (2) (y_1, y_2) is preferred to (x_1, x_2) . This is clearly absurd. In Figure 7.4 the consumer has apparently chosen (x_1, x_2) when she could have chosen (y_1, y_2) , indicating that (x_1, x_2) was preferred to (y_1, y_2) , but then she chose (y_1, y_2) when she could have chosen (x_1, x_2) —indicating the opposite!

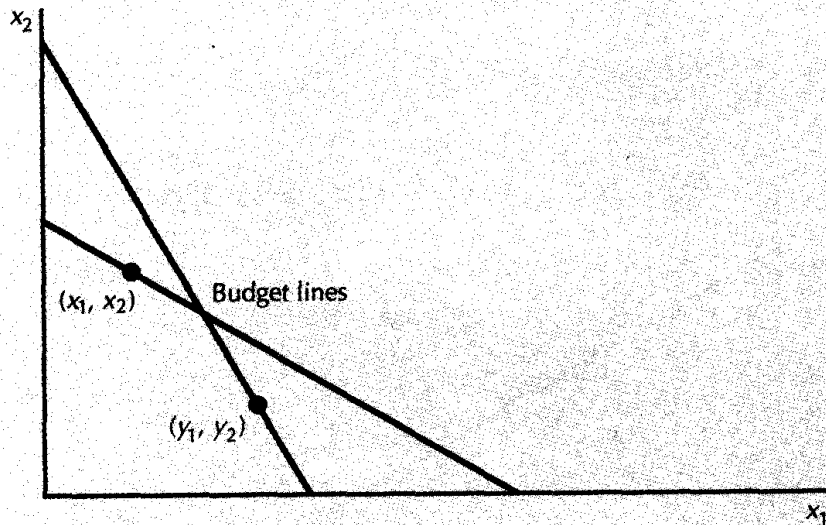
Clearly, this consumer cannot be a maximizing consumer. Either the consumer is not choosing the best bundle she can afford, or there is some other aspect of the choice problem that has changed that we have not observed. Perhaps the consumer's tastes or some other aspect of her economic environment have changed. In any event, a violation of this sort is not consistent with the model of consumer choice in an unchanged environment.

The theory of consumer choice implies that such observations will not occur. If the consumers are choosing the best things they can afford, then things that are affordable, but not chosen, must be worse than what is chosen. Economists have formulated this simple point in the following basic axiom of consumer theory

Weak Axiom of Revealed Preference (WARP). *If (x_1, x_2) is directly revealed preferred to (y_1, y_2) , and the two bundles are not the same, then it cannot happen that (y_1, y_2) is directly revealed preferred to (x_1, x_2) .*

In other words, if a bundle (x_1, x_2) is purchased at prices (p_1, p_2) and a different bundle (y_1, y_2) is purchased at prices (q_1, q_2) , then if

$$p_1x_1 + p_2x_2 \geq p_1y_1 + p_2y_2,$$



Violation of the Weak Axiom of Revealed Preference.
A consumer who chooses both (x_1, x_2) and (y_1, y_2) violates the Weak Axiom of Revealed Preference.

it must *not* be the case that

$$q_1 y_1 + q_2 y_2 \geq q_1 x_1 + q_2 x_2.$$

In English: if the y -bundle is affordable when the x -bundle is purchased, then when the y -bundle is purchased, the x -bundle must not be affordable.

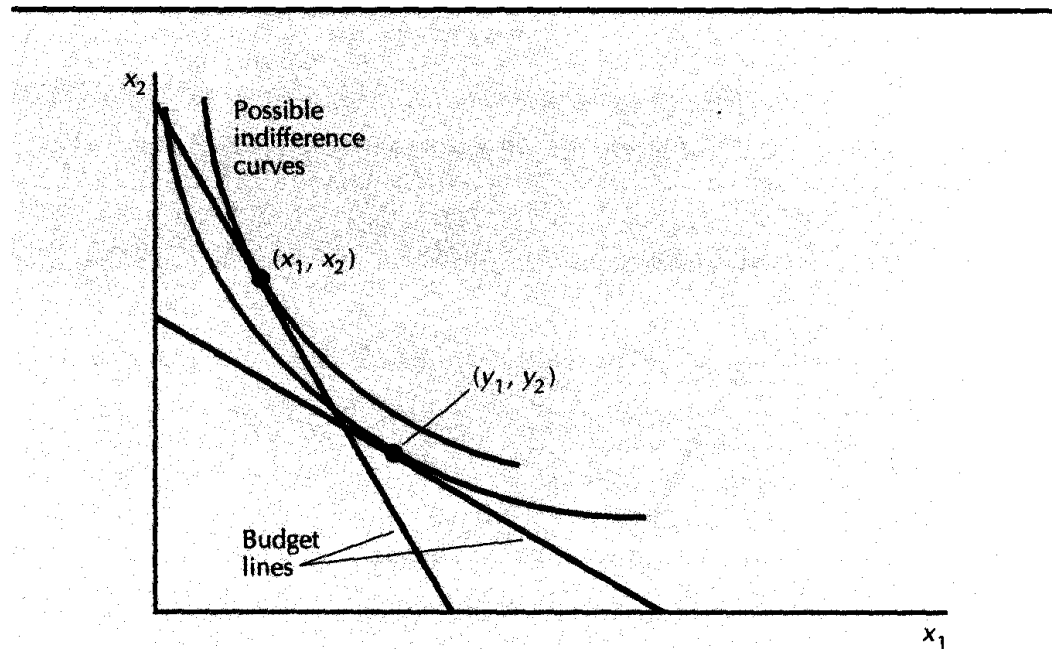
The consumer in Figure 7.4 has *violated* WARP. Thus we know that this consumer's behavior could not have been maximizing behavior.¹

There is no set of indifference curves that could be drawn in Figure 7.4 that could make both bundles maximizing bundles. On the other hand, the consumer in Figure 7.5 satisfies WARP. Here it is possible to find indifference curves for which his behavior is optimal behavior. One possible choice of indifference curves is illustrated.

tional 7.5 Checking WARP

It is important to understand that WARP is a condition that must be satisfied by a consumer who is always choosing the best things he or she can afford. The Weak Axiom of Revealed Preference is a logical implication

¹ Could we say his behavior is WARPed? Well, we could, but not in polite company.

Figure
7.5

Satisfying WARP. Consumer choices that satisfy the Weak Axiom of Revealed Preference and some possible indifference curves.

of that model and can therefore be used to check whether or not a particular consumer, or an economic entity that we might want to model as a consumer, is consistent with our economic model.

Let's consider how we would go about systematically testing WARP in practice. Suppose that we observe several choices of bundles of goods at different prices. Let us use (p_1^t, p_2^t) to denote the t^{th} observation of prices and (x_1^t, x_2^t) to denote the t^{th} observation of choices. To use a specific example, let's take the data in Table 7.1.

Table
7.1**Some consumption data.**

Observation	p_1	p_2	x_1	x_2
1	1	2	1	2
2	2	1	2	1
3	1	1	2	2

Given these data, we can compute how much it would cost the consumer to purchase each bundle of goods at each different set of prices, as we've

done in Table 7.2. For example, the entry in row 3, column 1, measures how much money the consumer would have to spend at the third set of prices to purchase the first bundle of goods.

Cost of each bundle at each set of prices.

		Bundles		
		1	2	3
Prices	1	5	4*	6
	2	4*	5	6
	3	3*	3*	4

The diagonal terms in Table 7.2 measure how much money the consumer is spending at each choice. The other entries in each row measure how much she would have spent if she had purchased a different bundle. Thus we can see whether bundle 3, say, is revealed preferred to bundle 1, by seeing if the entry in row 3, column 1 (how much the consumer would have to spend at the third set of prices to purchase the first bundle) is less than the entry in row 3, column 3 (how much the consumer actually spent at the third set of prices to purchase the third bundle). In this particular case, bundle 1 was affordable when bundle 3 was purchased, which means that bundle 3 is revealed preferred to bundle 1. Thus we put a star in row 3, column 1, of the table.

From a mathematical point of view, we simply put a star in the entry in row s , column t , if the number in that entry is less than the number in row s , column s .

We can use this table to check for violations of WARP. In this framework, a violation of WARP consists of two observations t and s such that row t , column s , contains a star *and* row s , column t , contains a star. For this would mean that the bundle purchased at s is revealed preferred to the bundle purchased at t and vice versa.

We can use a computer (or a research assistant) to check and see whether there are any pairs of observations like these in the observed choices. If there are, the choices are inconsistent with the economic theory of the consumer. Either the theory is wrong for this particular consumer, or something else has changed in the consumer's environment that we have not controlled for. Thus the Weak Axiom of Revealed Preference gives us an easily checkable condition for whether some observed choices are consistent with the economic theory of the consumer.

In Table 7.2, we observe that row 1, column 2, contains a star and row 2, column 1, contains a star. This means that observation 2 could have been

chosen when the consumer actually chose observation 1 and vice versa. This is a violation of the Weak Axiom of Revealed Preference. We can conclude that the data depicted in Tables 7.1 and 7.2 could not be generated by a consumer with stable preferences who was always choosing the best things he or she could afford.

7.6 The Strong Axiom of Revealed Preference

The Weak Axiom of Revealed Preference described in the last section gives us an observable condition that must be satisfied by all optimizing consumers. But there is a stronger condition that is sometimes useful.

We have already noted that if a bundle of goods X is revealed preferred to a bundle Y , and Y is in turn revealed preferred to a bundle Z , then X must in fact be preferred to Z . If the consumer has consistent preferences, then we should never observe a sequence of choices that would reveal that Z was preferred to X .

The Weak Axiom of Revealed Preference requires that if X is *directly* revealed preferred to Y , then we should never observe Y being *directly* revealed preferred to X . The **Strong Axiom of Revealed Preference (SARP)** requires that the same sort of condition hold for *indirect* revealed preference. More formally, we have the following.

Strong Axiom of Revealed Preference (SARP). *If (x_1, x_2) is revealed preferred to (y_1, y_2) (either directly or indirectly) and (y_1, y_2) is different from (x_1, x_2) , then (y_1, y_2) cannot be directly or indirectly revealed preferred to (x_1, x_2) .*

It is clear that if the observed behavior is optimizing behavior then it must satisfy the SARP. For if the consumer is optimizing and (x_1, x_2) is revealed preferred to (y_1, y_2) , either directly or indirectly, then we must have $(x_1, x_2) \succ (y_1, y_2)$. So having (x_1, x_2) revealed preferred to (y_1, y_2) and (y_1, y_2) revealed preferred to (x_1, x_2) would imply that $(x_1, x_2) \succ (y_1, y_2)$ and $(y_1, y_2) \succ (x_1, x_2)$, which is a contradiction. We can conclude that either the consumer must not be optimizing, or some other aspect of the consumer's environment—such as tastes, other prices, and so on—must have changed.

Roughly speaking, since the underlying preferences of the consumer must be transitive, it follows that the *revealed* preferences of the consumer must be transitive. Thus SARP is a *necessary* implication of optimizing behavior: if a consumer is always choosing the best things that he can afford, then his observed behavior must satisfy SARP. What is more surprising is that any behavior satisfying the Strong Axiom can be thought of as being generated by optimizing behavior in the following sense: if the observed choices satisfy SARP, we can always find nice, well-behaved preferences

that *could have* generated the observed choices. In this sense SARP is a *sufficient* condition for optimizing behavior: if the observed choices satisfy SARP, then it is always possible to find preferences for which the observed behavior is optimizing behavior. The proof of this claim is unfortunately beyond the scope of this book, but appreciation of its importance is not.

What it means is that SARP gives us *all* of the restrictions on behavior imposed by the model of the optimizing consumer. For if the observed choices satisfy SARP, we can “construct” preferences that could have generated these choices. Thus SARP is both a necessary and a sufficient condition for observed choices to be compatible with the economic model of consumer choice.

Does this prove that the constructed preferences actually generated the observed choices? Of course not. As with any scientific statement, we can only show that observed behavior is not inconsistent with the statement. We can’t prove that the economic model is correct; we can just determine the implications of that model and see if observed choices are consistent with those implications.

7.7 How to Check SARP

Let us suppose that we have a table like Table 7.2 that has a star in row t and column s if observation t is directly revealed preferred to observation s . How can we use this table to check SARP?

The easiest way is first to transform the table. An example is given in Table 7.3. This is a table just like Table 7.2, but it uses a different set of numbers. Here the stars indicate direct revealed preference. The star in parentheses will be explained below.

How to check SARP.

		Bundles		
		1	2	3
Prices	1	20	10*	22 ^(*)
	2	21	20	15*
	3	12	15	10

Now we systematically look through the entries of the table and see if there are any *chains* of observations that make some bundle indirectly revealed preferred to that one. For example, bundle 1 is directly revealed preferred to bundle 2 since there is a star in row 1, column 2. And bundle

2 is directly revealed preferred to bundle 3, since there is a star in row 2, column 3. Therefore bundle 1 is *indirectly* revealed preferred to bundle 3, and we indicate this by putting a star (in parentheses) in row 1, column 3.

In general, if we have many observations, we will have to look for chains of arbitrary length to see if one observation is indirectly revealed preferred to another. Although it may not be exactly obvious how to do this, it turns out that there are simple computer programs that can calculate the indirect revealed preference relation from the table describing the direct revealed preference relation. The computer can put a star in location st of the table if observation s is revealed preferred to observation t by any chain of other observations.

Once we have done this calculation, we can easily test for SARP. We just see if there is a situation where there is a star in row t , column s , *and* also a star in row s , column t . If so, we have found a situation where observation t is revealed preferred to observation s , either directly or indirectly, and, at the same time, observation s is revealed preferred to observation t . This is a violation of the Strong Axiom of Revealed Preference.

On the other hand, if we do not find such violations, then we know that the observations we have are consistent with the economic theory of the consumer. These observations could have been made by an optimizing consumer with well-behaved preferences. Thus we have a completely operational test for whether or not a particular consumer is acting in a way consistent with economic theory.

This is important, since we can model several kinds of economic units as behaving like consumers. Think, for example, of a household consisting of several people. Will its consumption choices maximize “household utility”? If we have some data on household consumption choices, we can use the Strong Axiom of Revealed Preference to see. Another economic unit that we might think of as acting like a consumer is a nonprofit organization like a hospital or a university. Do universities maximize a utility function in making their economic choices? If we have a list of the economic choices that a university makes when faced with different prices, we can, What it means is that SARP gives us a model of the consumer. For if the observed choices satisfy SARP, we can “construct” preferences that could have generated these choices. Thus SARP is both a necessary and a sufficient condition for the existence of a utility function that could have generated these choices. . . .

Suppose we examine the consumption bundles of a consumer at two different times and we want to compare how consumption has changed from one time to the other. Let b stand for the base period, and let t be some other time. How does “average” consumption in year t compare to consumption in the base period?

Suppose that at time t prices are (p_1^t, p_2^t) and that the consumer chooses (x_1^t, x_2^t) . In the base period b , the prices are (p_1^b, p_2^b) , and the consumer’s

choice is (x_1^b, x_2^b) . We want to ask how the “average” consumption of the consumer has changed.

If we let w_1 and w_2 be some “weights” that go into making an average, then we can look at the following kind of quantity index:

$$I_q = \frac{w_1 x_1^t + w_2 x_2^t}{w_1 x_1^b + w_2 x_2^b}.$$

If I_q is greater than 1, we can say that the “average” consumption has gone up in the movement from b to t ; if I_q is less than 1, we can say that the “average” consumption has gone down.

The question is, what do we use for the weights? A natural choice is to use the prices of the goods in question, since they measure in some sense the relative importance of the two goods. But there are two sets of prices here: which should we use?

If we use the base period prices for the weights, we have something called a **Laspeyres** index, and if we use the t period prices, we have something called a **Paasche** index. Both of these indices answer the question of what has happened to “average” consumption, but they just use different weights in the averaging process.

Substituting the t period prices for the weights, we see that the **Paasche quantity index** is given by

$$P_q = \frac{p_1^t x_1^t + p_2^t x_2^t}{p_1^t x_1^b + p_2^t x_2^b},$$

and substituting the b period prices shows that the **Laspeyres quantity index** is given by

$$L_q = \frac{p_1^b x_1^t + p_2^b x_2^t}{p_1^b x_1^b + p_2^b x_2^b}.$$

It turns out that the magnitude of the Laspeyres and Paasche indices can tell us something quite interesting about the consumer’s welfare. Suppose that we have a situation where the Paasche quantity index is greater than 1:

$$P_q = \frac{p_1^t x_1^t + p_2^t x_2^t}{p_1^t x_1^b + p_2^t x_2^b} > 1.$$

What can we conclude about how well-off the consumer is at time t as compared to his situation at time b ?

The answer is provided by revealed preference. Just cross multiply this inequality to give

$$p_1^t x_1^t + p_2^t x_2^t > p_1^t x_1^b + p_2^t x_2^b,$$

which immediately shows that the consumer must be better off at t than at b , since he could have consumed the b consumption bundle in the t situation but chose not to do so.

What if the Paasche index is *less* than 1? Then we would have

$$p_1^t x_1^t + p_2^t x_2^t < p_1^t x_1^b + p_2^t x_2^b,$$

which says that when the consumer chose bundle (x_1^t, x_2^t) , bundle (x_1^b, x_2^b) was not affordable. But that doesn't say anything about the consumer's ranking of the bundles. Just because something costs more than you can afford doesn't mean that you prefer it to what you're consuming now.

What about the Laspeyres index? It works in a similar way. Suppose that the Laspeyres index is *less* than 1:

$$L_q = \frac{p_1^b x_1^t + p_2^b x_2^t}{p_1^b x_1^b + p_2^b x_2^b} < 1.$$

Cross multiplying yields

$$p_1^b x_1^b + p_2^b x_2^b > p_1^b x_1^t + p_2^b x_2^t,$$

which says that (x_1^b, x_2^b) is revealed preferred to (x_1^t, x_2^t) . Thus the consumer is better off at time b than at time t .

7.9 Price Indices

Price indices work in much the same way. In general, a price index will be a weighted average of prices:

$$I_p = \frac{p_1^t w_1 + p_2^t w_2}{p_1^b w_1 + p_2^b w_2}.$$

In this case it is natural to choose the quantities as the weights for computing the averages. We get two different indices, depending on our choice of weights. If we choose the t period quantities for weights, we get the **Paasche price index**:

$$P_p = \frac{p_1^t x_1^t + p_2^t x_2^t}{p_1^b x_1^t + p_2^b x_2^t},$$

and if we choose the base period quantities we get the **Laspeyres price index**:

$$L_p = \frac{p_1^t x_1^b + p_2^t x_2^b}{p_1^b x_1^b + p_2^b x_2^b}.$$

Suppose that the Paasche price index is less than 1; what does revealed preference have to say about the welfare situation of the consumer in periods t and b ?

Revealed preference doesn't say anything at all. The problem is that there are now different prices in the numerator and in the denominator of the fractions defining the indices, so the revealed preference comparison can't be made.

Let's define a new index of the change in total expenditure by

$$M = \frac{p_1^t x_1^t + p_2^t x_2^t}{p_1^b x_1^b + p_2^b x_2^b}.$$

This is the ratio of total expenditure in period t to the total expenditure in period b .

Now suppose that you are told that the Paasche price index was greater than M . This means that

$$P_p = \frac{p_1^t x_1^t + p_2^t x_2^t}{p_1^b x_1^t + p_2^b x_2^t} > \frac{p_1^t x_1^t + p_2^t x_2^t}{p_1^b x_1^b + p_2^b x_2^b}.$$

Canceling the numerators from each side of this expression and cross multiplying, we have

$$p_1^b x_1^b + p_2^b x_2^b > p_1^b x_1^t + p_2^b x_2^t.$$

This statement says that the bundle chosen at year b is revealed preferred to the bundle chosen at year t . This analysis implies that if the Paasche price index is greater than the expenditure index, then the consumer must be better off in year b than in year t .

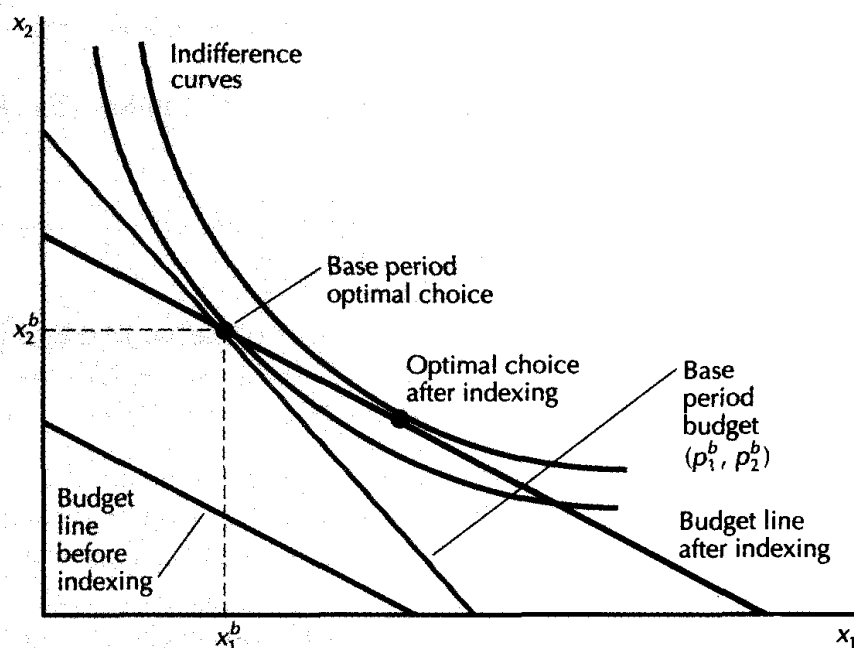
This is quite intuitive. After all, if prices rise by more than income rises in the movement from b to t , we would expect that would tend to make the consumer worse off. The revealed preference analysis given above confirms this intuition.

A similar statement can be made for the Laspeyres price index. If the Laspeyres price index is less than M , then the consumer must be better off in year t than in year b . Again, this simply confirms the intuitive idea that if prices rise less than income, the consumer would become better off. In the case of price indices, what matters is not whether the index is greater or less than 1, but whether it is greater or less than the expenditure index.

EXAMPLE: Indexing Social Security Payments

Many elderly people have Social Security payments as their sole source of income. Because of this, there have been attempts to adjust Social Security payments in a way that will keep purchasing power constant even when prices change. Since the amount of payments will then depend on the movement of some price index or cost-of-living index, this kind of scheme is referred to as **indexing**.

One indexing proposal goes as follows. In some base year b , economists measure the average consumption bundle of senior citizens. In each subsequent year the Social Security system adjusts payments so that the “purchasing power” of the average senior citizen remains constant in the sense that the average Social Security recipient is just able to afford the consumption bundle available in year b , as depicted in Figure 7.6.



Social Security. Changing prices will typically make the consumer better off than in the base year.

One curious result of this indexing scheme is that the average senior citizen will almost always be better off than he or she was in the base year b . Suppose that year b is chosen as the base year for the price index. Then the bundle (x_1^b, x_2^b) is the optimal bundle at the prices (p_1^b, p_2^b) . This means that the budget line at prices (p_1^b, p_2^b) must be tangent to the indifference curve through (x_1^b, x_2^b) .

Now suppose that prices change. To be specific, suppose that prices increase so that the budget line, in the absence of Social Security, would shift inward and tilt. The inward shift is due to the increase in prices; the tilt is due to the change in relative prices. The indexing program would then increase the Social Security payment so as to make the original bundle (x_1^b, x_2^b) affordable at the new prices. But this means that the budget line would cut the indifference curve, and there would be some other bundle

on the budget line that would be strictly preferred to (x_1^b, x_2^b) . Thus the consumer would typically be able to choose a better bundle than he or she chose in the base year.

Summary

1. If one bundle is chosen when another could have been chosen, we say that the first bundle is revealed preferred to the second.
2. If the consumer is always choosing the most preferred bundles he or she can afford, this means that the chosen bundles must be preferred to the bundles that were affordable but weren't chosen.
3. Observing the choices of consumers can allow us to "recover" or estimate the preferences that lie behind those choices. The more choices we observe, the more precisely we can estimate the underlying preferences that generated those choices.
4. The Weak Axiom of Revealed Preference (WARP) and the Strong Axiom of Revealed Preference (SARP) are necessary conditions that consumer choices have to obey if they are to be consistent with the economic model of optimizing choice.

REVIEW QUESTIONS

1. When prices are $(p_1, p_2) = (1, 2)$ a consumer demands $(x_1, x_2) = (1, 2)$, and when prices are $(q_1, q_2) = (2, 1)$ the consumer demands $(y_1, y_2) = (2, 1)$. Is this behavior consistent with the model of maximizing behavior?
2. When prices are $(p_1, p_2) = (2, 1)$ a consumer demands $(x_1, x_2) = (1, 2)$, and when prices are $(q_1, q_2) = (1, 2)$ the consumer demands $(y_1, y_2) = (2, 1)$. Is this behavior consistent with the model of maximizing behavior?
3. In the preceding exercise, which bundle is preferred by the consumer, the x-bundle or the y-bundle?
4. We saw that the Social Security adjustment for changing prices would typically make recipients at least as well-off as they were at the base year. What kind of price changes would leave them just as well-off, no matter what kind of preferences they had?
5. In the same framework as the above question, what kind of preferences would leave the consumer just as well-off as he was in the base year, for *all* price changes?